

**AUTHOR:**Martin Bekker¹ **AFFILIATION:**¹School of Electrical and Information Engineering, University of the Witwatersrand, Johannesburg, South Africa**CORRESPONDENCE TO:**

Martin Bekker

EMAIL:

martin.bekker@wits.ac.za

HOW TO CITE:Bekker M. Large language models and academic writing: Five tiers of engagement. *S Afr J Sci.* 2024;120(1/2), Art. #17147. <https://doi.org/10.17159/sajs.2024/17147>**ARTICLE INCLUDES:**

- ☐ Peer review
- ☐ Supplementary material

KEYWORDS:

large language models, academic integrity, academic writing, authorship, transparency

PUBLISHED:

30 January 2024



Large language models and academic writing: Five tiers of engagement

Significance:

Against a backdrop of the rapidly expanding use of large language models (LLMs) across diverse domains, this discussion breaks LLM usage into tiers of use, offering practical guidance to cautiously embrace the benefits of this significant new tool.

Introduction

2023 will be remembered as the year of large language models (LLMs), which, led by their brash poster child, ChatGPT, have changed the world forever. LLM-assisted writing will indelibly alter many writing tasks, offering speed and efficiency, and even automating-away many tasks.

However, academic, scientific, and intellectual integrity are at risk: not only due to mistakes that may creep in via the automation of writing, but also – and more importantly – owing to the loss of the ability to construct well-crafted arguments, ostensibly through the dulling of scholars' reasoning via the outsourcing of thinking that LLMs could engender. Moreover, ethical principles around intellectual process and ownership ought to be protected against the vague accountability of black-box algorithms (termed 'algorithmic opacity' by Eslami et al.¹) with respect to published or submitted work.^{2,3}

Academic journals, along with university and high-school curricula developers and assessment setters, need immediate yet thoughtful guidelines (rules and standards) for using LLMs and AI in the scientific process. In this Commentary, I propose a five-tier system that stipulates permissions and prohibitions around the use of LLMs in the academic writing process. I recapitulate what has changed, what – amid all the apparent changes – has stayed the same, and introduce the five tiers of LLM-assistance to academic writing, motivating the lines of distinction and suggesting appropriate uses.

What has changed?

In existence for over 40 years, language models are probabilistic models of a human language that can generate likelihoods of a series of words, based on text corpora on which they have been trained.⁴ Over the last decade, the size of the training text corpora and the number of weights between concepts held within the models have increased, necessitating affixing 'large' to recent models, now known as LLMs.⁵ ChatGPT was released in November 2022, combining the then-most advanced LLM with a chatbot-interface, simplifying the prompting (requesting) and serving (receiving responses) process. With their promise of speed and efficiency, ChatGPT and other LLMs have had an immediate impact on the academe; demonstrating the ability to automate the writing of reports, research and literature papers, exams, and computer code, among others, to various degrees of human ability, with each iteration showing improvement.

Contemporary LLMs represent a break with the past, based on (1) the speed and scale of information processing, (2) an unprecedented function of research process assistance (which includes research summary and data manufacture), and (3) the potential for the outsourcing of thought. The first of these three innovations often accompanies new technological tools. However, the scale and speed at which LLMs perform information processing tasks have now surpassed the human performance of certain tasks within the so-called information economy⁶, suggesting a quantum leap in functioning and utility.

The second change is tied to the very nature of the way LLMs can process information. The ability of transformer models – the deep-learning architecture behind LLMs – to manipulate text (or language, or indeed anything that can be represented as a language) has made LLMs superlative at summarising texts, style transfers (ranging from translation to mere tone tweaks) and spelling and grammar corrections. Moreover, in addition to LLMs, several other AI-related tools that assist in the research process have recently been introduced (with many more to follow in their wake). With these, two particular functions spring to mind: first, research summarising tools (e.g. Elicit, Perplexity and Consensus) that can 'find' work (published but unknown to the scholar), 'understand' discourses, identify research gaps, and assist with literature reviews. The second are those that can manufacture artificial (or synthetic) data. Here, a scholar might give instructions regarding what a data set should contain, and, in the absence of this being available (as secondary data) or impossible to gather (for, say, ethical reasons), such data can be 'created' instantly.

The third change relates to the potential for the wholesale contracting out, or 'outsourcing' of thought to a (non-human) algorithm. While cheating is nothing new – academics have long been aware of student essays drafted by 'essay mills', computer code written by friends or copied from online repositories, or even a human double sitting for an exam on behalf of a less-prepared student – LLMs present the academe with a new level of concern. That is, from students to scientists, writers are now able to turn to LLMs to author academic work from conceptualisation through research and writing⁷, putting in peril the principle of scientific advancement through human reasoning (or sound thinking as articulated through writing).



Perils notwithstanding, the immediate benefit of such a tool – one that can fix the register, grammar, and punctuation of a text in seconds and apparently at no cost to the user – is immediate and clear. This is especially the case for those writing in a second language, which case applies to the majority of academic scholars, who must publish in English. This ‘levelling of the playing field’ is to be welcomed by the academic community.^{2,3,8} Academics often recite that good writing is indistinguishable from good thinking; the corollary of this is that clear, high-quality writing helps not only non-native speakers get the recognition they deserve, but benefits humanity if readers access knowledge in clear, correct, and accessible language.

The immediate drawback accompanying this new class of technology is that, sadly, many things that the academy has long battled to counter – dishonesty, cheating and plagiarism – have almost instantly become much harder to detect, owing to the increased volume and sophistication of the breaches. Moreover, LLMs are known to routinely produce credible untruths (‘hallucinations’, ‘simulated authority’⁸ or ‘compelling misinformation’⁹) and omit attributions of their source or training data (plagiarism). Combined with the outsourcing of thought itself, such concerns render the use of LLMs for academic work potentially disingenuous at best, and at worst, in violation of the norms of scientific research. With all of this in mind, the need for practical guidance through the ethical minefield of LLMs is clear.

What has remained constant?

It is comforting that, despite the impressive and daunting changes wrought by the advent of LLMs, in reality, most scientific principles endure. First, the three values of beneficence, autonomy, and justice, all tied to non-maleficence and the avoidance of suffering¹⁰ stand firm; in this sense, right is still right, and wrong remains wrong. Similarly, cheating and plagiarism remain anathema to the spirit of science, while openness, reproducibility, and the sharing (non-obscuring or gatekeeping) of data, keeping in mind all the caveats of harm, still stand as ideals.

Also holding firm are peer review as a well-established principle before publication, and the less-formal review-by-peers; that is, the shaping and improving of ideas (and writing) based on conversations, correspondence, arguments and counsel. Technical help – whether in the form of word processors, spell checkers, pocket calculators and software programs, or human editors and proofreaders – remains accepted and welcomed.

The five-tier system: An ethical guide to using LLMs

The present scramble to incorporate LLM use in academic work (or to find ways to ban it) implies that conversations about ethical guidelines and AI-use standards are timeous and valuable. Given the promise and peril of the new, but guided by the three values (justice, autonomy and beneficence), I propose a five-tier system to simplify thinking around permissions and prohibitions related to using LLMs for academic writing. While representing increasing ‘levels’ of LLM support that progress along a seeming continuum, the tiers in fact represent paradigmatically different types of mental undertakings.

Tier 1: Use ban

The first level comprises a complete ban on LLM-based support. This means that no LLM tools may be used in the preparation of the academic text. This tier therefore implies the highest level of human authorship and research authenticity.

Given its Draconian nature, coupled with the likelihood of inadvertent violations (e.g. the spelling and grammar checks employed by ‘ordinary’ word processors use a form of AI, and common word processors will soon be incorporating several other dimensions of LLMs), this tier is the most inviting to be flouted. Difficulty to enforce, lack of benefit and abundance of drawbacks (e.g. a step backwards in terms of present practice, given, for example, the ubiquity of automatic spelling and grammar checks) make such a tier likely only to be used in very specific circumstances, such as proctored university examinations or other formal testing conditions.

Tier 2: Proofing tool

Here, human-written text may be submitted to an LLM, accompanied by a prompt instructing the model to fix spelling, grammar, register, tone, and style (in the manner that products such as Grammarly might do). A proofing tool can be instructed to catch (and recommend remedies for) tone and style variations, identify problematic or misused words, and highlight direct translations.

The point here, of course, is that the work is presented at the end of the writing process – once the experimental and argumentative thinking is complete. Tier 2 does not outsource the thinking (or pain) that goes into the drafting process; rather, it takes fleshed-out thoughts and cosmetically enhances (or translates) them in much the same manner as would a private or in-house proofing team (or academic translation service).

Facilitating word-perfect (or near enough) text before submission, this level of usage can ‘level the playing field’ for non-native speakers and early-career researchers; that is, allowing those authors to display their arguments and findings to best advantage, smoothing over linguistic objections (tacit or implicit), and thus allowing work to be judged on merit alone. Of course, it may also enable ‘lazy’ writing (perhaps on the part of native English writers, knowing that sloppy writing will be fixed by an algorithm). Nonetheless, clarity and universality are significant in matters of standards, and thus, under this tier, any author might make use of this LLM proofing assistance.

Tier 3: Copyediting tool

Editing and proofing are distinguished by their sequential place in the preparation of texts, and also in the tasks that they perform: this distinction is echoed in the differences between Tiers 2 and 3. Given Tier 3 permissions, an author may ask for an LLM to alter text beyond correcting linguistic mistakes and aligning stylistic requirements. Shortening wordy text (e.g. reducing an abstract from 300 to 150 words), expanding for clarity, and rephrasing for precision are tasks every academic writer has laboured over and would likely welcome assistance with. Other areas of assistance at this tier would be checking citations for accuracy, style, and appropriateness – all things an LLM editorial function can do, and, in most cases, that paid in-house editors do to ensure the quality of publications (see Table 1 for prompt examples).

An editorial function can ensure a text is well organised, making changes to a text’s structure by reordering paragraphs, or highlighting missing arguments. This tier will also include the assistance that Tier 2 permits,

Table 1: Examples of initial prompts per Tiers 2 and 3 (these are intentionally kept straightforward and unsophisticated)

Tier	Sample prompts
2	a) The following is section [x] from my paper, which I aim to submit to [name of journal]. Proofread the section and suggest corrections for spelling, grammar, tone, and style errors. Ensure the text is clear and free from any direct translations. Also, review the section for tone and style variations, and note any such pointwise. Finally, identify problematic or misused words, list each, and provide a recommendation for replacement.
	b) Make recommendations to improve the overall readability and coherence of the text. “[text here]”
3	a) Edit the below to shorten to 2000 words while preserving content, intention and clarity.
	b) Check the citations in the following document for accuracy, style, and appropriateness. List any necessary corrections pointwise.
	c) Make suggestions for the reorganising of paragraphs in the document below to improve the overall structure and flow of the argument. Briefly motivate each modification.



including spelling, grammar, tone and style corrections, flagging unusual words, nonsensical or confusing text, and guiding smooth transitions between paragraphs.

This tier would best be used during and after the writing process, and would likely be used iteratively, perhaps once a substantial section has been written.

Efficient editing (which can be a tedious and expensive journey) and the clarity of resultant texts are among the chief benefits of this tier. If correctly applied, critical thinking (and laboratory work or experimentation) would have preceded this step that in principle simply allows for a near-flawless write-up. Nonetheless, as with Tier 2, intellectual thoroughness and writing rigour may be compromised, ostensibly making this a compromise that must be accepted. This tier raises the point that copyediting is, and should be, regarded as an intellectual contribution (although copyeditors are seldom credited in scholarly journal articles in the way they may be in the publication of books), underscoring the observation that perhaps all forms of support ought to be acknowledged in academic writing.

Tier 4: Drafting consultant

Tier 4 speaks to a process of human–LLM ‘co-creation’ (‘augmented writing’¹⁷ or ‘coauthoring’²²). In addition to the support permitted under Tiers 2 and 3, this tier permits an iterative back-and-forth of ideas as one might do with a coauthor, up to the point of (and including) the LLM suggesting the omission of certain arguments, suggesting alternative ‘interpretations’, or requesting that one rerun experiments or check back to confirm previous findings.

In this tier, the author can interact with an LLM to plan a research write-up and shape and develop an argument, including requesting sample lines (e.g. instructing the LLM to ‘compose an opening line’). Thus, this is not merely a tool offering a ‘substantive edit’, but a tool that can ensure one’s evidence backs up one’s argument (that is, where an LLM might even contribute to shaping that argument at earlier stages), and, where this is not the case, can provide suggestions on how to remedy such gaps.

Unlike the previous tiers, Tier 4 implies LLM engagement before the writing process commences, followed by iterative ‘reporting back’ sessions with the LLM as the writing advances. Potential benefits include the support such a routine would lend to early-career researchers: ‘handholding’ and sense-checking their writing process and arguments, while also offering suggestions and criticisms throughout the process to ensure quality. While the approach allowed by this tier is likely to significantly speed up the writing process, it does appear to tip the scales in terms of potential risks. Hallucinations and biases (subtle or not), both artefacts of LLMs, are more likely to manifest in co-created works. Also, LLMs may establish ways to obscure poor research behind apparently brilliant writing. It also follows that an overreliance on LLMs – here operating well beyond argumentative and stylistic considerations – would

constitute a loss of skills, the mastery of which would be expected in professional scientists or academics.

Tier 5: No limits

The fifth tier allows any LLM assistance at any stage. This tier includes brainstorming avenues of research, discussing and suggesting hypotheses and ideas, and even allowing the LLM to write text on the author’s behalf. Such a no-holds-barred approach also includes interpreting results, summarising other scholarly work, and suggesting the implication of findings. In other words, Tier 5 permits the outsourcing of thought.

This tier is likely to be useful for instructing students on AI literacy and AI usage, and more generally demonstrating the dangers of stochastic models. This tier has limited value to scholarly, peer-reviewed publications, as the principle of authorship and the requisite of originality (at least as currently conceived) would likely be violated by, for example, a systematic review conducted solely by an AI agent.

Overview of tiers

With their growing levels of permissions, the tiers represent not only increasing degrees of LLM support, but also increasing levels of LLM dependence. Table 2 illustrates the tiers, and typical use-case points of entry, alongside the most obvious advantages and disadvantages of each.

As the tiers ‘progress’, so do the apparent speed and efficiency of tasks, as well as the dangers of LLM hallucination and manipulation – both significant and sensible concerns, given the opaque nature of LLM’s stochastic processes and governance, not to mention the monopolising tendencies among Big Tech in general. Added to this is the arguably less immediately pernicious loss of academic integrity that would attend the outsourcing of thinking, and may come to represent a threat to humanity’s overall ability to undertake quality scientific research, in the unlikely scenario in which LLM reliance is left unchecked and unregulated. Moreover, real and unconscionable human exploitation¹¹ and high environmental costs¹², both present in current LLM models, cannot be discounted or wished away.

Other AI tools available to the would-be academic author

While not part of the writing process, and thus adjacent to the proposed tiers, other academic tools draw on new advances in LLM or other recent machine-learning technologies. The ethical and intellectual considerations of their use overlap with those of LLMs, as do the efficiencies they offer researchers and writers.

The first is a class of **research summarising tools** already introduced in this paper. These can perform scans of literature on a topic (e.g. Perplexity), provide one-line summations of research papers (e.g. Elicit), give the apparent scientific consensus on a topic (e.g. Consensus),

Table 2: Summary of large language model (LLM) permission tiers, where they come into the writing process, and their most obvious benefits and risks

Tier	Effect / type of tool	Place in the writing process	Most obvious benefit	Most obvious risk
1	Ban	n/a	Ensures 100% human authorship and does not compromise academic integrity	Inevitably flouted except under exam conditions
2	Proofing	After	Increases efficiency, reduces cost	Might subtly alter meaning or obscure intentions
3	Editing	During	Produces well-organised, word-perfect writing	Language may be bland, may foster laziness
4	Co-creating	Start	Offers an alternative to human partnership–making interpretative and instructive suggestions, error checks	Authorship is opaque, high risk of introducing hallucinations and biases, inexplicability, loss of autonomy, loss of critical reasoning, outsourcing of thought
5	All	All	Enables high-speed academic writing, maximum support for inexperienced academics	As per Tier 4, but more extreme

or even create a literature review ‘at the touch of a button’. Such tools can provide excellent assistance in exploring new fields of research or fields related to one’s own work (see Jansen et al.¹³ for a discussion of areas in which LLMs might support survey-based research), but are not substitutes for the process of critically reading to inform and order one’s own intuitions and conclusions, gradually bringing new ideas into relation with one’s own thought-scape. Moreover, such models may play a role in reifying conventional wisdoms, and in so doing, drown out marginal voices (the latter which may also be thought of as ‘majority voices’, considering that most people, including academics, are non-Western, non-white, and non-anglophone, despite the outsized influence of Western universities on the global scientific community).

The second non-writing LLM-based tool is ‘**automatic data analysis**’ (e.g. Langchain), whereby data sets can be loaded up to an LLM for statistical analysis. In one dimension, the use of such technologies is equivalent to that of a pocket calculator: a logical time-saver, provided the user understands what the LLM is doing. For example, for at least the past three decades, scientists have routinely used multiple regressions, typically executed by statistical software or a coding routine in a software library. Use of statistical software (not least the writing up of results) requires a basic knowledge of statistics and data analysis. Subject to new developments, danger enters when the process of statistical analysis is not understood by the scientist, but regarded, crudely, as magic (i.e. it cannot be explained). In all, the use of LLM interfaces for statistical analysis will likely become commonplace, to the benefit of science in general, with the qualification that scientists will still be required to understand at least ‘the bare bones’ of statistical analysis.

The third application of AI tools, now in the form of machine learning, is the **creation of synthetic data** (e.g. Stalice). Here, one can ask for data containing certain characteristics and of a given (potentially vast) size, or create data for teaching or illustrative purposes (including instances of data unavailability or cases where ethical or legal considerations proscribe the gathering of such data, such as sensitive healthcare data with patient identifiers). Such data sets will play an increasing role in teaching and testing, and so long as users keep in mind that the data in hand are fabricated, will be of great advantage in several domains (and are currently used in the training of self-driving cars, models for financial service security, and the pharmaceutical industries).

It is reasonable to expect that more AI-based tools will join the arsenal of the scientist. While caution (based, as ever, on the beneficence, autonomy, and fairness and the avoidance of maleficence principles) regarding the tools’ creation, application and implications must be exercised, many of these tools will provide important and progressive support to scientific advancements⁹, and should be embraced.

LLM tools and safety principles

Returning to LLM tools and how they can support the academic writing process, two inviolable principles merit further reflection: ownership and transparency. AI raises complex questions about these two principles; however, for the time being, the five-tier system, plus the appropriate use of supplementary material, may help to clarify questions around authorship, responsibility, and where we should stand on the place of the thinking process.

Ownership is the tenet that a submitted or published work and all of its contents remain the responsibility of a human author, and that the author is the only accountable party for mistakes or other consequences emanating from the work.³ Any academic will be familiar with examples of authors choosing to omit their names from a publication (despite having contributed to the scholarship) when they do not fully agree with the contents, or feel unable to be held responsible for arguments contained in the work. Yet the broader point here is that the owner of ideas, the agent bearing the risk, and the agent deciding to be listed as author of a work is, so far, a human one. In our scientific pursuit, in ‘advancing on Chaos and the Dark’ it is the human – often individual – thinker who toils, who weighs, who risks.¹⁴ Intellectual progress and the advancement of human thinking assumes the hitherto generally unspoken assumption that human thought ought not to be outsourced to non-human entities. Differently put, and evoking Chesterton¹⁵, you

cannot make science without a soul. Even in the scenario in which a human and an LLM ‘co-create’ a work, the responsibility for the content still must rest somewhere, and, in the spirit of science, this would most obviously be the human author. Similarly, a recent US court ruling heard that ‘the guiding human hand’ is a ‘bedrock requirement’ to authorship.¹⁶ The idea of blaming an LLM for mistakes appears disingenuously evasive and points to a concerning ambiguity over authorship. Certainly, under the first three tiers, as suggested here, authorship, and therefore ownership, rests with the human writer. The corollary of this is that the scientific value of papers, books and artefacts produced under Tiers 4 and 5 are de facto limited, and are likely to remain so: should this change, we will have to rethink authorship and academic credit anew.

While it is possible that humans will not forever be regarded as the sole culpable parties of their publications, recent work drawing on AI advances continues to confirm the principle of human authorship: despite the success of AlphaFold being based on Google DeepMind’s algorithms, the celebrated *Nature* paper¹⁷ lists only the human authors. Of course, acknowledgement of AI tools used in research (and acknowledgements in general) is categorically different from named authors of a work; that is, contribution does not imply attribution.

The second principle is that of **transparency**. Transparency – that is, showing one’s work and thought process – lies at the heart of scientific accountability, reproducibility, peer accessibility, and public trust.² LLMs are by their nature opaque^{1,5}; that does not mean they cannot be used, but rather that we must be open about *when* and *how* we use them. The alternative scenario is that readers have to guess whether LLMs have been used in the production of text³ – here, mistrust is a solvent of credibility.

I suggest that authors who use LLM assistance include a way for their readers to see the prompts (and responses) they used in preparing a text. While the American Psychological Association referencing standard has issued a protocol for referencing ChatGPT, possibly a better way of referencing LLM contributions would be to provide a way for readers to access the entire series of human–LLM interactions, including every prompt and response. A hyperlink to an online repository (such as GitHub or Google documents) may risk being too fleeting a base, while a text file as supplementary material might suffice in ‘showing one’s working’ in the same way as sharing one’s routine of commands to statistical software, or sharing a codebase when programming.

Discussion: AI hype and despair

There is nothing emergent – in the complex adaptive systems theory sense – in the working of a pocket calculator. Calculators’ answers are consistent, predictable, replicable, and regular. In this sense, LLMs are not like calculators: their massive size and black-box nature appear to have given them, at least by most accounts, emergent capabilities. The general public (and technical!) discourse has tended to label this not as ‘emergent’ but rather as ‘generative’ – which is, all told, succumbing to the language of AI hype. 2023 may well be marked as a high point of AI optimism, not least owing to the abilities and potential of LLMs. However, the ability to distinguish helpful innovation from unhelpful hyperbole (whether AI saviourism or AI catastrophising) is important in order to keep humankind’s present problems and struggles in perspective, recognise our immediate moral duties, and rationally analyse the extent to which a new class of tools can help (or hinder) scientific progress and human betterment.

For scientists across every branch of knowledge, mental panic and ossification remain our nemeses. We gain most by seeing neither cataclysmic doom nor total redemption in technology, but instead recalibrating a new technology’s value based on what it can change, and what it can’t.

Among the proposed tiers, Tier 1 is techno-pessimistic. This tier assumes that technology per se represents a threat to knowledge production and human capabilities. This position is, to my mind, untenable for an academic journal already heavily reliant on LLMs (e.g. in the form of spell-checkers), not to mention calculator-like technology. In contrast, Tiers 2 and 3 may be regarded as technology-embracing. Optimistic about the

efficiencies LLMs bring to human knowledge and scientific advancement, these tiers advocate for adoption, remove the first-language barriers so often inhibiting the global dissemination of great ideas, and may even expedite the writing-up process. If permitted some rumination, one might suggest that Tiers 4 and 5 are leaning towards AI hype: both of these supposing that LLMs are either on the path to true cognitive supremacy and should thus be employed at all costs (the slave will soon become the benign master), or, alternatively, taking up the despairing position that LLMs will soon be so ubiquitous that any resistance to their use is bound to fail. One might argue that King² and Jansen et al.¹³, on whose insights this note draws, lean towards the possibility of a Tier 5 future, where AI will become a colleague and coauthor.

Conclusion

Five tiers of LLM support for academic writing have been introduced, each offering a different level of writing support, and each entering the writing (and thought) process at a different stage. With some intentionality, the principles of ownership (plus responsibility) and transparency (sharing of prompts) can, and ought to be maintained.

Competing interests

I have no competing interests to declare.

References

1. Eslami M, Vaccaro K, Lee MK, Elazari Bar On A, Gilbert E, Karahalios K. User attitudes towards algorithmic opacity and transparency in online reviewing platforms. In: Proceedings of the 2019 CHI Conference on Human Factors in Computing Systems; 2019 May 4–9; Glasgow, Scotland. New York: Association for Computing Machinery; 2019. p. 1–14. <https://doi.org/10.1145/3290605.3300724>
2. King MR. A place for large language models in scientific publishing, apart from credited authorship. *Cell Mol Bioeng*. 2023;16:95–98. <https://doi.org/10.1007/s12195-023-00765-z>
3. Li H, Moon JT, Purkayastha S, Celi LA, Trivedi H, Gichoya JW. Ethics of large language models in medicine and medical research. *Lancet Digit Health*. 2023;5(6):e333–e335. [https://doi.org/10.1016/S2589-7500\(23\)00083-3](https://doi.org/10.1016/S2589-7500(23)00083-3)
4. Rosenfeld R. Two decades of statistical language modeling: Where do we go from here? *Proc IEEE*. 2000;88(8):1270–1278. <https://doi.org/10.1109/5.880083>
5. Bender EM, Gebru T, McMillan-Major A, Shmitchell S. On the dangers of stochastic parrots: Can language models be too big? In: Proceedings of the 2021 ACM Conference on Fairness, Accountability, and Transparency; 2021 March 3–10. New York: Association for Computing Machinery; 2021. p. 610–623. <https://doi.org/10.1145/3442188.3445922>
6. Bubeck S, Chandrasekaran V, Eldan R, Gehrke J, Horvitz E, Kamar E, et al. Sparks of artificial general intelligence: Early experiments with GPT-4 [preprint]. *arXiv:arXiv:2303.12712*; 2023. <https://doi.org/10.48550/arXiv.2303.12712>
7. Kulesz O. The impact of large language models on the publishing sectors: Books, academic journals, newspapers [master's thesis]. Sweden: Linnaeus University; 2023. Available from: <https://urn.kb.se/resolve?urn=urn:nbn:se:lnu:diva-119366>
8. Rillig MC, Ågerstrand M, Bi M, Gould KA, Sauerland U. Risks and benefits of large language models for the environment. *Environ Sci Technol*. 2023;57(9):3464–3466. <https://doi.org/10.1021/acs.est.3c01106>
9. Spitale G, Biller-Andorno N, Germani F. AI model GPT-3 (dis) informs us better than humans [preprint]. *arXiv:arXiv:2301.11924*; 2023. <https://doi.org/10.1126/sciadv.adh1850>
10. US National Commission for the Protection of Human Subjects of Biomedical and Behavioral Research. The Belmont report: Ethical principles and guidelines for the protection of human subjects of research. Washington DC: US Department of Health and Human Services; 1979. Available from: <http://www.hhs.gov/ohrp/regulations-and-policy/belmont-report>
11. Perrigo B. Exclusive: OpenAI used Kenyan workers on less than \$2 per hour to make ChatGPT less toxic. *Time*. 18 January 2023. Available from: <https://time.com/6247678/openai-chatgpt-kenya-workers/>
12. Strubell E, Ganesh A, McCallum A. Energy and policy considerations for deep learning in NLP. In: Korhonen A, Traum D, Marquez L, editors. Proceedings of the 57th Annual Meeting of the Association for Computational Linguistics; 2019 July 28 – August 02; Florence, Italy. Kerrville, TX: Association for Computational Linguistics. p. 3645–3650. <https://doi.org/10.18653/v1/P19-1355>
13. Jansen BJ, Jung SG, Salminen J. Employing large language models in survey research. *Nat Lang Process J*. 2023;4:100020. <https://doi.org/10.1016/j.nl.2023.100020>
14. Emerson RW. Self-Reliance (1841). In: Porte J, editor. *Essays and lectures*. New York: Library of America; 1983. p. 261.
15. Chesterton GK. *Orthodoxy*. San Francisco, CA: Ignatius Press/John Lane Company; 1908.
16. Thaler Sv, Perlmutter S, Register of Copyrights and Director of the U.S. Copyright Office, et al., Civil Action No. 22-1564 (BAH) (18.08.2023). US District Court for the District of Columbia; 2023.
17. Jumper J, Evans R, Pritzel A, Green T, Figurnov M, Ronneberger O, et al. Highly accurate protein structure prediction with AlphaFold. *Nature*. 2021;596(7873):583–589. <https://doi.org/10.1038/s41586-021-03819-2>